

TeleTrust-Konferenz

Berlin, 29.06.2023

KI im Recht: Anforderungen an Datenschutz und IT-Sicherheit

Bernhard Kloos, HK2 Comtection

Bernhard Kloos



- Partner/ Rechtsanwalt, HK2 Rechtsanwälte
- Geschäftsführer HK2 Comtection GmbH
- Autor zu Datenschutz in KI & Recht kompakt

Handelsblatt**Deutschlands
BESTE
Anwälte****2022****Bernhard Kloos**
IT-RechtHandelsblatt • 24.06.2022
Eine Kooperation mit**Best Lawyers**

Agenda

- Rechtsrahmen für KI
- Maßnahmen und Risiken
- Datenschutzrechtliche Anforderungen



Vorschlag einer EU-Verordnung zur Regulierung Künstlicher Intelligenz vom 21.04.2021 (KI-Verordnung)

HK2

Cometion

- Verordnungsentwurf zur Förderung sicherer und rechtskonformer KI im gesamten Binnenmarkt
- Einheitlicher Rechtsrahmen: Harmonisierte und unmittelbar wirkende Vorschriften für die Entwicklung, das Inverkehrbringen, die Inbetriebnahme und die (bestimmungsgemäße) Verwendung künstlicher Intelligenz in der Europäischen Union
- KI soll allen nutzen und dem Gemeinwohl dienen, Regeln sollen klar und robust sein
- EU will Führungsrolle übernehmen
- Status: Einigung des EU-Parlaments auf einen [Vorschlag am 14.06.2023](#); Trilog-Verhandlungen mit den EU-Mitgliedsstaaten und der Kommission
- Abschluss 2023/2024 erwartet



KI-Verordnung: Grundsätze

- **Definition KI-System:** ein **maschinengestütztes System**, das so konzipiert ist, dass es mit unterschiedlichem Grad an **Autonomie** operieren kann und das für explizite oder implizite Ziele **Ergebnisse** wie Vorhersagen, Empfehlungen oder Entscheidungen hervorbringen kann, **die das** physische oder virtuelle **Umfeld beeinflussen** (Art. 3 (1) Nr. 1 = OECD Definition)
- **Zweck:**
 - Förderung **menschenzentrierter und vertrauenswürdiger künstlicher Intelligenz**
 - Schutz der **Gesundheit, der Sicherheit, der Grundrechte, der Demokratie und der Rechtsstaatlichkeit sowie der Umwelt**
- **Datenschutz und andere Rechtsakte** (z.B. E-Privacy, Produktsicherheit, Verbraucherschutz) **bleiben unberührt**, d.h. gelten **parallel**
- **Ausnahmen** etwa für Militär, Wissenschaftliche Forschung und Entwicklung, Behörden von vertrauenswürdigen Drittstaaten im Rahmen internationaler Übereinkünfte



KI-Verordnung: Struktur

- **Risikobasierter Ansatz**

- **unannehmbares Risiko / Verbot**

- z.B. Manipulation ohne Einwilligung mit spürbarer Beeinträchtigung und erheblichem Schaden, biometrische Kategorisierung oder Echtzeit-Fernidentifikation in öffentlich zugänglichen Räumen, Predictive Policing, Ableitung von Emotionen in den Bereichen Strafverfolgung, Grenzmanagement, Arbeitsplatz und Bildung

- **Hochrisikosysteme**

- z.B. Recruiting, Personalmanagement, KRITIS, Überwachung bei Prüfungen, Sozialhilfe, Kreditwürdigkeit, Kranken- oder Lebensversicherung, Justiz, Empfehlungssystem bei sehr großen Social-Media-Plattformen,
 - wenn sie ein erhebliches Risiko für die Gesundheit, die Sicherheit oder die Grundrechte von natürlichen Personen darstellen -> Eindeutige Leitlinien sollen folgen

- **geringes oder minimales Risiko**

- z.B. Chatbots, Spamfilter -> ggf. Transparenzpflichten



KI-Verordnung – Anforderungen an Hochrisiko-KI

Beispiele I

HK2

Comtection

- **Risikomanagementsystem**
 - Analyse, Bewertung, Maßnahmen während des Lebenszyklus der KI -> Risiken müssen nach vernünftigem Ermessen vertretbar sein
 - Risikomanagementsystem kann in bereits bestehende Risikomanagementverfahren integriert werden.
 - Inhalte: Ermittlung, Abschätzung und Bewertung der bekannten und vernünftigerweise vorhersehbaren Risiken, Ergreifung geeigneter und gezielter Risikomanagementmaßnahmen zur Bewältigung
- **Obligatorisches Testen** in jedem Fall vor dem Inverkehrbringen oder der Inbetriebnahme für die Zweckbestimmung und vernünftigerweise vorhersehbare Fehlanwendung
- **Data Governance:** Verwendung qualitätsgesicherter Trainings-, Validierungs- und Testdatensätze (einschließlich auf Relevanz, Verfügbarkeit, Verzerrungen / Bias und Lücken oder Mängel)



KI-Verordnung – Anforderungen an Hochrisiko-KI

Beispiele II

HK2

Comtection

- **Dokumentation und Aufzeichnung** / Protokollierung gemäß dem Stand der Technik -> enormes Haftungspotential
- **Transparente Informationen und Gebrauchsanweisungen in allen Sprachen** -> Anbieter und Nutzer müssen Funktionsweise hinreichend verstehen können
- **Menschliche Aufsicht** im angemessenen Verhältnis zu den Risiken **und „Stopptaste“**
- **Genauigkeit, Robustheit** (gegen Fehler, Störungen und Unstimmigkeiten, mit Angabe des Genauigkeitsgrades), **Sicherheit und Cybersicherheit**
- **Optimaler Schutz:** technische und organisatorische Maßnahmen, um sicherzustellen, dass Hochrisiko-KI-Systeme so widerstandsfähig wie möglich gegenüber Fehlern, Störungen oder Unstimmigkeiten sind, die innerhalb des Systems oder der Umgebung, in der das System betrieben wird, insbesondere wegen seiner Interaktion mit natürlichen Personen oder anderen Systemen auftreten können



KI-Verordnung – Anforderungen an Hochrisiko-KI

Beispiele III

HK2

Cometion

- technische Lösungen für den **Umgang mit KI-spezifischen Schwachstellen** zum Schutz vor
 - Datenvergiftung
 - Modellvergiftung
 - feindliche Beispiele oder Modellvermeidung
 - Angriffe auf vertrauliche Daten oder Modellmängel, die zu einer schädlichen Entscheidungsfindung führen könnten
- **Qualitätsmanagementsystem** (einschließlich Beobachtung, Meldung von Vorfällen und Rechenschaft / persönliche Verantwortung) -> auf der Grundlage von Normen wie ISO 9001 kein doppeltes Qualitätsmanagementsystem eingerichtet werden, sondern eine Anpassung der vorgenommen werden
- **Konformitätsbewertungsverfahren** bei Produkten wie Maschinen, Spielzeugen, Aufzügen, Funkanlagen, Druckgeräten, Sportbootausrüstung, Seilbahnen, Geräte, Medizinprodukte



Sicherheit von KI-Systemen – Angriffstypen (BSI)

Bundesamt für Sicherheit und Informationstechnik (BSI) hat mehrere Studien auf dem Weg zu Prüfkriterien für KI-Systeme in Auftrag gegeben. Klassische Angriffstypen:

- **Backdoor Attacks**

Während des Trainings werden Schwachstellen in das KI-System eingefügt, die dann vom Angreifer ausgelöst werden können.

- **Poisoning Attacks/ Vergiftungsangriffe:**

Die Vergiftungsproben werden während des Trainings der KI injiziert und führen zu einer Beeinträchtigung

- **Evasion/Umgehungsangriffe:**

Durch Manipulation von Eingabedaten verleiten Angreifer das KI-Modell im Betrieb zu vom Entwickler nicht vorgesehenen Ausgaben.

- **Data & Model Extraction Attacks:**

Angreifer extrahieren Informationen hinsichtlich der Trainingsdaten bis hin zum kompletten Trainingsmuster aus dem angewandten Modell oder die Funktionalität des Modells.

Sicherheit von KI-Systemen – Threats and Vulnerabilities (ENISA)

- **ENISA** hat mehrere Veröffentlichungen zu Risiken und deren Behandlung veröffentlicht, etwa:
 - AI CYBERSECURITY CHALLENGES – Threat Landscape (2020)
 - SECURING MACHINE LEARNING ALGORITHMS (2021)
 - STUDY: CYBERSECURITY AND PRIVACY IN AI – MEDICAL IMAGING DIAGNOSIS (2023)
- Breiterer Ansatz, auch Datenschutz und interne Risiken, etwa
 - Menschliche Fehler
 - Unrechtmäßige oder unfaire Verarbeitung
 - Verstoß gegen Grundsatz der Datenminimierung
- Enthalten sind z.B. Vorschläge für generic and specific Cybersecurity and Privacy Controls



NEFARIOUS ACTIVITY/ABUSE

- Unauthorized access to data sets and data transfer process
- Manipulation of data sets and data transfer process
- Unauthorized access to models models' code
- Compromise and limit AI results
- Compromising AI inference's correctness data
- Compromising ML inference's correctness algorithms
- Data poisoning
- Data tampering
- Elevation of Privilege
- Insider threat
- Manipulation of optimization algorithm
- Misclassification based on adversarial examples
- Model poisoning
- Transferability of adversarial attacks
- Online system manipulation
- Model Sabotage
- Scarce data
- White box , targeted or non targeted
- Introduction of selection bias
- Manipulation of labelled data
- Backdoor insert attacks on training datasets
- Overloading confusing labelled dataset
- Compromising ML training validation data
- Compromising ML training augmented data
- Adversarial examples
- Reducing data accuracy
- ML Model integrity manipulation
- ML model confidentiality
- Compromise of data brokers providers
- Manipulation of model tuning
- Sabotage
- DDoS
- Access Control List ACL) manipulation
- Compromising ML pre processing
- Compromise of model frameworks
- Corruption of data indexes
- Reduce effectiveness of AI ML results
- Label manipulation or weak labelling
- Model backdoors



PHYSICAL ATTACK

- Errors or timely restrictions due to non reliable data infrastructures
- Model Sabotage
- Infrastructure system physical attacks
- Communication networks tampering
- Sabotage



DISASTER

- Natural disasters (earthquake , flood, fire , etc)
- Environmental phenomena heating , cooling , climate change



FAILURES/MALFUNCTIONS

- Compromising AI application viability
- Errors or timely restrictions due to non reliable data infrastructures
- 3 rd party provider failure
- ML Model Performance Degradation
- Scarce data
- Stream interruption
- Inadequate absent data quality checks
- Lack of documentation
- Weak requirements analysis
- Poor resource planning
- Weak data governance policies
- Compromising ML pre processing
- Corruption of data indexes
- Label manipulation or weak labelling
- Compromise of model frameworks



EAVESDROPPING/INTERCEPTION/HIJACKING

- Data inference
- Data theft
- Model Disclosure
- Stream interruption
- Weak encryption



LEGAL

- Corruption of data indexes
- Compromise privacy during data operations
- Profiling of end users
- Lack of data protection compliance of 3 rd parties
- Vendor lock in
- SLA breach
- Weak requirements analysis
- Lack of data governance policies
- Disclosure of personal information
- Corruption of data indexes



OUTAGES

- Infrastructure/system outages
- Communication networks outages



UNINTENTIONAL DAMAGES/ACCIDENTAL

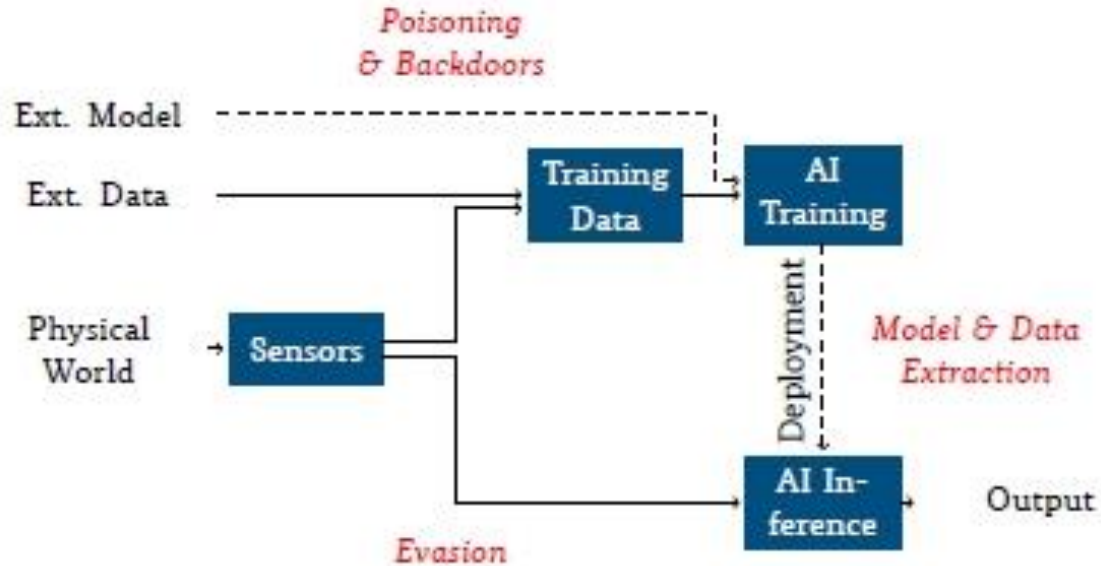
- Compromise and limit AI results
- Compromise privacy during data operations
- Compromising AI inference's correctness data
- Compromising feature selection
- Compromising ML inference's correctness algorithms
- Misconfiguration or mishandling of AI system
- ML Model Performance Degradation
- Online system manipulation
- Lack of sufficient representation in data
- Mishandling of statistical data
- Manipulation of labelled data
- Compromising ML training augmented data
- Reducing data accuracy
- Compromise of data brokers providers
- Erroneous configuration of models
- Bias introduced by data owners
- Label manipulation or weak labelling
- Disclosure of personal information
- Compromise of model frameworks

Quelle: ENISA, AI CYBERSECURITY CHALLENGES, <https://www.enisa.europa.eu/publications/artificial-intelligence-cybersecurity-challenges/@download/fullReport>, S. 29

Sicherheit von KI-Systemen (BSI)

HK2

Cometion



Risikomatrix

Quelle: https://www.bsi.bund.de/SharedDocs/Downloads/EN/BSI/KI/Security-of-AI-systems_fundamentals.pdf

S. 4

HK2

Cometcon

Risk	Low	Medium	High
Deployment Environment	Known lab environment	Confined open-world deployment	Publicly available online service
Training Data Origins	Well established data sets	Self-created data sets	Data sets of unknown origin
Inference Data Origins	Local, e.g., a fixed data set	Air gap, e.g., a camera system	External queries, e.g., the internet
Degree of Autonomy	Recommendation to human expert	Decision observed by human expert	Autonomous decision by the AI
Output Type	System decision	Discretized model output	Full model output
User Access Origins	Local admins	All internal users	Global access
User Access	Single query	Limited number of queries	Unlimited queries
Model Origins	Locally trained model with custom architecture	Locally retrained public architecture	Publicly available model
Learning Strategy	Single training	Continuous and regular updates	Continuous and spontaneous updates

Sicherheit von KI-Systemen (BSI)

HK2

Comtection

Unterschiedliche Kenntnisstände des Angreifers sind zu betrachten:

Beobachtung	Teilkenntnis	Vollständige Kenntnis
KI-Entscheidung, Output Verteilung	Architektur, Trainingsaufbau	Verteidigungen, Architektur, Trainingsaufbau
Black-Box	Gray-Box	White-Box

KI-Verordnung – Anforderungen an Anbieter eines Basismodells (Art. 28 b)

HK2

Comtection

- Aufstieg von ChatGPT führte zur **Einführung des Basismodells**, worunter die „Generative KI“ fällt
 - **Basismodell** = breite Datenbasis, allgemeine Ausgabe, breite Palette unterschiedlicher Aufgaben
 - Abgrenzung zur „**einfachen**“ **Mehrzweck-KI-Systemen**
 - „**Generative KI**“ = Basismodell + dazu bestimmt, mit unterschiedlichem Grad an Autonomie Inhalte wie komplexe Texte, Bilder, Audio- oder Videodateien zu generieren
- Pflichten gelten für:
 - **Eigenständiges Modell**
 - **eingebettetes Modell** in KI-System, Produkt oder unter Open-Source-Lizenzen (z.B. ChatLLaMA)
- Einzuhalten sind **vergleichbare Pflichten wie bei Hochrisikosystemen**, d.h. Risikomanagement, Data-Governance, sichergestelltes Niveau an Leistung, Vorhersagbarkeit, Interpretierbarkeit, Korrigierbarkeit, Sicherheit und Cybersicherheit, (Bereitstellung) der technische Dokumentation für 10 Jahre, Qualitätsmanagementsystem
- Bei „**generativer KI**“ **gelten zusätzliche Pflichten** (Art. 28 b Abs. 4 KI-VO)
 - Zusätzliche Transparenz
 - Gestaltung und Weiterentwicklung zur Einhaltung des Rechts nach Stand der Technik
 - Veröffentlichung welche urheberrechtlich geschützte Daten für Training verwendet wurden



Einrichtung eine Internationalen KI- Aufsicht

- KI als Superintelligenz braucht eine spezielle Behandlung und Koordination
 - *“We must mitigate the risks of today’s AI technology too, but superintelligence will require special treatment and coordination.” (Open AI-Team)*
- KI wird grenzüberschreitend entwickelt
- Internationale Atomenergiebehörde (IAEO) als Vorbild
- KI-Verordnung greift das bereits auf mit dem Einrichtung des Europäischen Amts für künstliche Intelligenz (Art. 56)
- Eindämmung der potenziellen Gefahr bzgl. der internationalen Verbreitung von Fehlinformationen
- Bislang zersplitterte aufsichtsbehördliche Aktivitäten zu ChatGPT z.B. in Italien, Kanada, USA, Frankreich, Spanien, EDSA; Deutschland (Hessen, Baden-Württemberg und Rheinland-Pfalz)

Explainable AI als ein zentraler Faktor

- **Nachvollziehbarkeit** als wichtigstes Forschungsziel der KI, sonst ist menschliche Kontrolle nicht möglich
- **Nachvollziehbarkeit beruht auf Transparenz und Erklärbarkeit**
- **Erklärbar $\neq$ Komplettbeweis** (z.B. bei Black Box oder neuronalen Netzwerken)
- **Erklärungsstrategie** abhängig von Zielgruppen, verwendeten Datentypen und KI-Modell
- **Modell- und / oder Entscheidungserklärung**
- **Erklärungswerkzeuge** (z.B. LIME, SHAP, Saliency Maps, Surrogat-Modelle) müssen mindestens genau so schnell weiter entwickelt werden wie KI-Systeme
- Weg zu einer einfachen Generalklausel? **KI muss „nur“ rechtmäßig (d.h. auch diskriminierungsfrei), richtig, risikoarm und erklärbar sein.**

Datenschutzrechtliche Eckpunkte (gemäß DSK)

- Für KI-Systeme gelten die **Grundsätze für die Verarbeitung personenbezogener Daten, etwa Transparenz und Zweckbindung** (Art. 5 DSGVO)
- Grundsätze müssen durch frühzeitig geplante **technische und organisatorische Maßnahmen** (TOM) von Verantwortlichen umgesetzt werden (Art. 25 DSGVO)
- Die **Sicherheit der Verarbeitung muss gewährleistet sein** unter Berücksichtigung des **Standes der Technik**, der Implementierungskosten und der Art, des Umfangs, der Umstände und der Zwecke der Verarbeitung sowie der unterschiedlichen Eintrittswahrscheinlichkeit und Schwere des Risikos (Art. 32 DSGVO)
- **Gewährleistungsziele** sind **Transparenz und Erklärbarkeit, Datenminimierung, Nichtverkettung, Intervenierbarkeit, Verfügbarkeit, Integrität und Vertraulichkeit**

Rechtsgrundlagen

- Einwilligung, d.h.
 - freiwillig ohne Druck oder Einfluss
 - spezifisch
 - informiert (Zweck, Funktion, Dauer, Dritte)
 - aktiv und unzweideutig
 - jederzeit widerruflich
- Vertragsdurchführung
- Berechtigtes Interesse
- Rechtliche Pflicht



Technische und organisatorische Maßnahmen (TOM) bei Entwicklung und Betrieb von KI-Systemen

(gemäß DSK

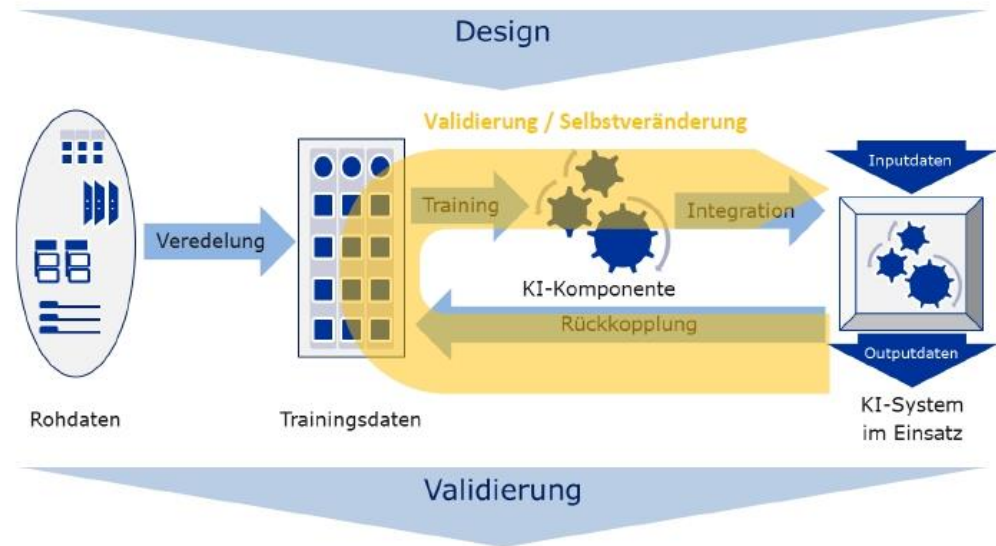
https://www.datenschutzkonferenz-online.de/media/en/20191106_positionspapier_kuenstliche_intelligenz.pdf, Grafik S. 2)

HK2

Comtection

Grundlage ist der allgemeine Lebenszyklus von KI-Systemen:

1. **Design** des KI-Systems und deren KI-Komponenten
2. **Veredelung der Rohdaten** zu Trainingsdaten
3. **Training** der KI-Komponenten
4. **Validierung** der Daten und KI-Komponenten sowie angemessene Prüfungsmethoden
5. **Nutzung** des KI-Systems
6. **Rückkopplung** von Ergebnissen und Selbstveränderung des Systems



TOM Design

Festzulegen und zu dokumentieren sind:

- **Zweck** des KI-Systems und der KI-Komponenten
- angestrebte **Güte** des KI-Systems
- **Eingabe- und Ausgabeparameter**
- **Hypothese des gewünschten Zusammenhangs** zwischen Eingabe- und Ausgabeparametern
- **Festlegung von unerwünschtem Verhalten**
- Festlegung eines geeigneten KI-Verfahrens und Dokumentation dessen Auswahl
- Beteiligte Personen und Rechte

TOM Veredelung

In dieser Phase sind zu dokumentieren:

- **Herkunft der Rohdaten** klären, normalisieren, standardisieren, komplettieren, fehlerbereinigen und möglichst reduzieren und / oder pseudonymisieren
- Festlegung des **Verfahrens zur Datenveredelung**
- **Ausschluss** diskriminierender oder ungewünschter Einflussfaktoren (**Bias**)
- **Relevanz und Repräsentativität der Trainingsdaten**
- Möglichkeiten zur Anonymisierung oder Synthetisierung
- Abschätzung der Datenmengen, um angestrebte Güte zu erreichen
- **Erkennung und Korrektur** von falschen oder diskriminierenden Roh- und Trainingsdaten, Testdaten und Testverfahren
- **Manipulationsfreiheit, Vertraulichkeit und Verfügbarkeit von Roh- und Trainingsdaten**

TOM Training

- Prüfung der **Kompatibilität der Zwecke** der KI-Komponente mit den Verwendungszwecken der Trainingsdaten
- **Dokumentation der Trainingsdaten** und Klärung ihrer Herkunft
- Festlegung von Primär-Trainings-, Verifikations- und Testdatensätzen einschließlich des Verfahrens zur Einteilung und den Regeln zur Verwendung
- **Verfahren zur Ermittlung der Güte** einschließlich ggf. Prozess zur Erweiterung der Trainingsdaten, um die Güte zu erreichen
- **Verhinderung von unbefugten Manipulationen** an KI-Komponenten

TOM Validierung

- **Test mit gewünschten und mit störsignalbehafteten (synthetischen) Testdaten** auf Basis des erwarteten und unerwünschten Verhaltens
- **Prüfung der Hypothese und der Erwartungen**
- **Dokumentation der Güte** der KI-Komponente inklusive der ermittelten Fehlerrate und Systemstabilität,
- Untersuchung der KI-Komponente auf **Erklärbarkeit und Nachvollziehbarkeit**
- **Evaluation** des ausgewählten KI-Verfahrens bezüglich alternativer, erklärbarer KI-Verfahren

TOM Einsatz I

- **Überwachung des Verhaltens** der KI-Komponente und des -Systems
- Möglichkeit des **Eingreifens einer Person** in den Entscheidungsprozess
- Entscheidungen mit hohen Risiken dürfen von KI-Systemen nur vorbereitet werden, Möglichkeit des Stoppens oder des Ersetzens einer KI-Komponente durch Fall-Back-Lösung
- **Protokollierung von finalen Entscheidungen**, deren Freigabe/Bestätigung/Ablehnung, Zeitpunkt und ggf. entscheidende Person
- **Prüfung und Bewertung** der Hypothese und der Erwartungen auf Basis des Verhaltens im Betrieb sowie der Güte des KI-Systems und seiner KI-Komponenten

TOM Einsatz II & Rückkopplung

- regelmäßige **Evaluierung welche Eingabe- und Ausgabeparameter** der KI-Komponenten für das gewünschte Verhalten des KI-Systems relevant und erforderlich sind, ggf. Anpassung
- Berechnungen der KI-Komponente möglichst auf Endgerät des Nutzer ohne Übermittlung der Daten (**Federated Learning**), ansonsten auf Geräten des Verantwortlichen
- bei Übermittlung der Daten an KI-Komponente **Ende-zu-Ende-Verschlüsselung** einsetzen
- **Auskunftsmöglichkeit für Betroffene** zum Zustandekommen von Entscheidungen, Prognosen sowie über das Verfahren zur Ermittlung der Güte des KI-Systems und die festgestellte Güte
- **Ausschluss der Nutzung durch Unbefugte** und zu nicht vorgesehenen Zwecken (Berechtigungskonzept)
- **Mechanismen zum Anfechten und Korrigieren** vorsehen und darauf hinweisen
- **Verhinderung von unbefugten Manipulationen**
- **Rückkopplung von erzeugten Outputdaten** zum weiteren Training, möglichst anonymisiert oder pseudonymisiert

TOM Datenschutz

Generell: Die **Maßnahmen** sind **so zu gestalten**, dass der Verantwortliche seinen **datenschutzrechtlichen Pflichten in allen Verarbeitungsphasen und -schritten nachkommen kann**, etwa

- Information (z.B. Hinweisschild auf Sensoren beim autonomen Fahren)
- Auskunft
- Berichtigung
- Löschung
- Einschränkung der Verarbeitung
- Widerspruchsrecht
- Widerruf der Einwilligung
- Verbot automatisierter Entscheidung

In a nutshell

- **Wissen was reinkommt** (Data Governance)
- **Datenschutzrechtliches Konzept**
 - Synthetische Daten, Anonymisierung, Pseudonymisierung
 - Rechtsgrundlage passend zum Verwendungszweck
 - Risikoschonende Struktur (fördert, verschlüsselt oder (erst einmal) mit pseudonymisierten oder unsensiblen Daten)
 - Datentöpfe getrennt halten
- **Verstehen, was rauskommt**: Erklärbarkeit und Nachvollziehbarkeit, Ausschluss von Bias
- **Risikomanagement** von Beginn und durch alle Phasen implementieren
 - Welche gibt es?
 - Wie gravierend sind sie?
 - Maßnahmen zur Eindämmung
- **Testen** auf erwartete Ergebnisse, auch mit falschen Testdaten und mithilfe der Community
- **Überwachung und menschliche Aufsicht**
- **Dokumentation** (einschließlich QS-Maßnahmen, Zuständigkeiten) und Protokollierung
- **Transparente Information**



Haben Sie Fragen?

HK2

Comtection



www.it-sicherheit-und-recht.de

www.hk2.eu/newsletter

www.comtection.de